

Estudo Piloto de Detecção Ultrassônica de Etanol em Misturas Água–Etanol: Utilização de Algoritmos de Inteligência Artificial

1st Mateus F. Silva
PPGEE UFOP-UNIFEI
Universidade Federal de Itabubá
Itabira, MG, Brasil
mateusfilipi22@unifei.edu.br

2nd Maxwell Damasceno
PPGEE UFOP-UNIFEI
Universidade Federal de Itabubá
Itabira, MG, Brasil
maxwell@unifei.edu.br

3rd Marcelo M. Tiago
PPGEE UFOP-UNIFEI
Universidade Federal de Ouro Preto
João Monlevade, MG, Brasil
marcelomtiago@ufop.edu.br

4th Ricardo T. Higuti
Departamento de Engenharia Elétrica
Universidade Estadual Paulista
Ilha Solteira, SP, Brasil
ricardo.t.higuti@unesp.br

5th Luiz C. B. Torres
PPGEE UFOP-UNIFEI
Universidade Federal de Ouro Preto
João Monlevade, MG, Brasil
luiz.torres@ufop.edu.br

Resumo—Este trabalho propõe um método de regressão para estimar a concentração de etanol em misturas aquosas (fração mássica [w/w]) a partir de medições ultrassônicas. Indicadores de energia, potência e variação foram extraídos das séries temporais de ultrassom e usados como entrada para três algoritmos de aprendizado de máquina: *Random Forest*, *Gradient Boosting* e *Deep Neural Network*. A avaliação considerou RMSE, MAE e coeficiente de determinação (R^2). Em conjuntos de 24 amostras, o *Gradient Boosting* obteve os melhores resultados (RMSE = 0,1715; MAE = 0,1218; R^2 = 0,693), seguido pelo *Random Forest*. A *Deep Neural Network* apresentou maior sensibilidade ao baixo volume de dados. Os resultados obtidos mostram que o modelo de *Boosting* é mais estável e se comporta melhor nos diferentes cenários de treinamentos. Embora o erro médio ainda supere o limite de 0,7% definido por normas do setor, os resultados iniciais demonstram a viabilidade de abordagens baseadas em *ensemble* para predição de concentrações de etanol em aplicações industriais.

Palavras-chave—Ultrassom; Aprendizagem de Máquina; Misturas de Água e Etanol.

I. INTRODUÇÃO

Desenvolver métodos acessíveis e precisos para mensurar grandezas físicas representa um desafio significativo na indústria. Instrumentos que capturam com precisão variáveis críticas de processos químicos e físicos são essenciais para assegurar a qualidade e a uniformidade dos produtos [1]. Em especial, no setor petroquímico, a determinação rápida e precisa do teor de etanol em misturas aquosas ([w/w]) é fundamental em cadeias produtivas de biocombustíveis, farmacêuticas e alimentícias, pois variações na concentração podem afetar desempenho de equipamentos, gerar emissões indesejadas e comprometer a conformidade com normas ambientais e de qualidade [2] [3].

Métodos tradicionais de análise de concentração, como HPLC, refratometria e alcoometria, oferecem elevada pre-

cisão em laboratório, mas demandam tempo, calibrações frequentes e infraestrutura especializada, limitando sua aplicação em sistemas de automação e em tempo real [4]. Modelos termodinâmicos como COSMO-SAC LinSandler2002 e NRTL são amplamente usados para prever comportamento de soluções hidroalcoólicas, porém dependem de parâmetros experimentais e podem falhar em condições dinâmicas de processo.

Abordagens recentes têm demonstrado a eficácia do ultrassom na estimativa de concentração de etanol de forma não invasiva. [5] utilizou a velocidade do som em função da temperatura para prever a concentração de etanol em misturas, obtendo alta precisão com interpolação bilinear. Já [6] propôs o uso de redes neurais artificiais para prever a densidade de misturas etanol-água, superando modelos termodinâmicos tradicionais. Por fim, [7] combinou a velocidade ultrassônica com regressão PLS e SVR para estimar etanol e metanol em bebidas comerciais. Apesar dos bons resultados, tais abordagens não exploram diretamente características temporais do sinal bruto do transdutor utilizado.

Neste trabalho, propõe-se a caracterização de misturas de água e etanol por meio de sensores ultrassônicos, explorando vários indicadores extraídos das séries temporais do sensor, como estatística descritiva, potência dos sinais, análise espectral, largura de banda, área dos sinais, taxa de decaimento e velocidade de propagação. Esses parâmetros são utilizados como entradas para três modelos de aprendizado de máquina: *Random Forest*, *Gradient Boosting* e *Deep Neural Network*. As amostras foram preparadas conforme a norma OIML para concentração em fração mássica [W/W], e os algoritmos foram avaliados por meio de métricas de regressão (RMSE, MAE e R^2).

Este estudo piloto é conduzido sobre um volume reduzido de dados em um *dataset* ainda pouco explorado, cujo principal desafio é atingir resultados robustos com amostras limitadas. Diversas iniciativas em áreas como ciências médicas [8] [9], astrofísica [10] [11], ecologia [12] e indústria [13]

O presente trabalho foi realizado com apoio da Coordenação de Aperfeiçoamento de Pessoal de Nível Superior – Brasil (CAPES) – Código de Financiamento 001, do Conselho Nacional de Desenvolvimento Científico e Tecnológico (CNPQ).

demonstram a viabilidade de abordagens específicas para mitigar a escassez de dados e otimizar a eficácia dos modelos.

O trabalho está organizado da seguinte forma. Na Seção II, definimos objetivos e descrevemos o preparo de amostras e extração de indicadores; na Seção III apresentamos a metodologia dos modelos e configuração de treinamento; na Seção IV detalhamos resultados, comparações com literatura e discussão crítica; e na Seção V concluímos destacando ganhos, limitações e propostas de trabalhos futuros.

II. METODOLOGIA

A. Sistema ultrassônico

A Figura 1 apresenta um diagrama esquemático da célula de medição ultrassônica utilizada neste trabalho.

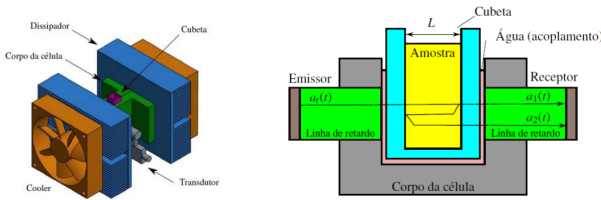


Fig. 1. Esquema do sistema de medição ultrassônico utilizado.

O sistema é composto por um bloco metálico, com temperatura controlada por meio de células de efeito Peltier. Um medidor de temperatura de precisão (Isotech, modelo Millik) foi utilizado como referência para controlar a temperatura das amostras analisadas, com um erro de regime permanente menor do que $0,01^{\circ}\text{C}$. Cubetas de quartzo, com caminho de propagação de 10 mm, foram utilizadas para depositar as amostras de interesse.

Dois transdutores de 75 MHz (Olympus, modelo V2022), operando em modo transmissão-recepção, foram utilizados para emitir e receber os sinais acústicos. Uma camada de água destilada foi utilizada para garantir o acoplamento entre transdutores e cubeta. Os transdutores foram excitados por meio de um pulsador (JSR Ultrasonics, modelo DPR500), com um pulso estreito de -400 V .

Os sinais elétricos provenientes do transdutor receptor foram adquiridos por meio de uma placa de aquisição (Signatec, modelo PX14400), com frequência de amostragem de 400 MHz e resolução vertical de 14 bits. O *software* Matlab foi utilizado para controlar os sistemas de excitação e aquisição. Mais detalhes a respeito do sistema de aquisição utilizado podem ser obtidos em [14].

B. Preparação das amostras

As amostras analisadas neste trabalho foram preparadas utilizando água destilada e etanol HPLC. Uma balança analítica, com resolução de 0,1 mg, foi utilizada para definir as proporções de água e etanol em cada mistura. Foram preparadas amostras com concentrações variando entre 10% [w/w] e 90% [w/w], utilizando frascos Schott de 250 ml.

As misturas foram agitadas por 10 min e, em seguida, mantidas em repouso por 24 h. Após esse período, um densímetro de bancada (Anton Paar, modelo DMA 4500) foi utilizado para medir a densidade das amostras. As tabelas

de acoultimetria, definidas por OIML [15], foram utilizadas para converter os valores de densidade medidos em valores de concentração [w/w] de etanol em água.

C. Pré-processamento e seleção de características

Os dados adquiridos abrangem 24 amostras rotuladas, totalizando 108 observações. Para cada concentração específica de água e álcool, foram realizadas entre três e cinco aquisições do sinal. As características extraídas incluem a amplitude e a duração do sinal emitido, além do primeiro sinal refletido (eco) na amostra. A Figura 2 ilustra um exemplo de amostra com 20% de etanol.

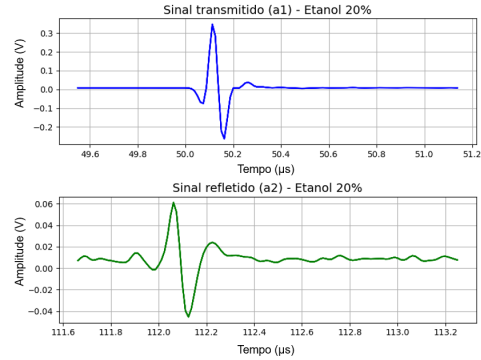


Fig. 2. Representação dos sinais adquiridos no domínio do tempo.

Para cada nível de concentração, realizamos múltiplas medições ultrassônicas a fim de reduzir ruídos e garantir a robustez dos dados. No entanto, o preparo de cada amostra é um procedimento laborioso e de custo elevado, e as replicatas coletadas pertencem à mesma batelada, ou seja, compartilham exatamente o mesmo valor de concentração. Assim, optamos por selecionar uma única observação representativa para cada rótulo, resultando em 24 amostras distintas. Essa escolha preserva a variabilidade inerente ao processo de preparação — as múltiplas coletas asseguram que o sinal capturado seja confiável — sem inflar artificialmente o conjunto de treinamento com repetições de um mesmo valor de referência. Dessa forma, garantimos que cada amostra contribua de maneira única para o aprendizado dos modelos, sem enviesar as métricas de avaliação.

Para realizar a tarefa de predição da concentração, os sinais foram processados e resumidos em indicadores, utilizados como entradas para os modelos. Esses indicadores foram calculados a partir dos sinais acústicos mencionados anteriormente.

Os indicadores utilizados incluem média, desvio padrão, variância, potência dos sinais obtida através da média do quadrado das amplitudes, frequência dominante (determinada pela Transformada Rápida de Fourier (FFT)), largura de banda (calculada com base nas frequências que correspondem à queda de potência em -3 dB), a integral da curva que representa a área sob os sinais, a taxa de decaimento (estimada a partir da amplitude ao longo do tempo), e a velocidade de propagação do som no meio (calculada com base no tempo necessário para o sinal acústico se propagar pela amostra). A Tabela I apresenta as expressões utilizadas para calcular esses indicadores.

TABLE I
EXPRESSÕES UTILIZADAS PARA CALCULAR OS INDICADORES.

Indicador	Fórmula
Média	$\mu = \frac{1}{N} \sum_{i=1}^N x_i$
Desvio Padrão	$\sigma = \sqrt{\frac{1}{N} \sum_{i=1}^N (x_i - \mu)^2}$
Variância	$\sigma^2 = \frac{1}{N} \sum_{i=1}^N (x_i - \mu)^2$
Máximo Absoluto	$\max = \max(x_i)$
Potência do Sinal	$P = \frac{1}{T} \int_0^T s(t) ^2 dt$
Frequência Dominante	$f_d = \arg \max FFT(s(t)) $
Largura de Banda	$BW = f_{high} - f_{low}$
Integral da Curva	$A = \int_0^T s(t) dt$
Taxa de Decaimento	$D(t) = D_0 e^{-\frac{t}{\tau}}$
Velocidade de Propagação	$v = \frac{d}{\Delta t}$

Os indicadores foram inicialmente regularizados para garantir média zero e desvio padrão unitário. Em seguida, a técnica de Análise de Componentes Principais (PCA) foi utilizada para reduzir o número de características utilizadas pelo modelo, mantendo apenas as informações mais relevantes. No entanto, ao longo do desenvolvimento, observou-se que o uso do PCA prejudicava o desempenho dos modelos. Assim, essa etapa foi posteriormente descartada.

D. Definição dos modelos

As arquiteturas de modelos analisadas neste trabalho incluem *Gradient Boosting*, *Random Forest* e *Deep Neural Network*, cada uma com suas particularidades e aplicações. O *Gradient Boosting* é um método de aprendizado de máquina que combina múltiplos modelos (geralmente árvores de decisão) para criar um modelo robusto, focando iterativamente nas previsões incorretas dos modelos anteriores [16]. Por outro lado, o *Random Forest* é uma técnica de *ensemble* que utiliza várias árvores de decisão, construídas a partir de subconjuntos aleatórios dos dados, para melhorar a precisão e reduzir o risco de *overfitting* [17]. As *Deep Neural Networks* (DNNs) são estruturas complexas de redes neurais que consistem em múltiplas camadas, permitindo a modelagem de relações não lineares complexas nos dados [18].

Neste estudo, para garantir um desempenho otimizado dos modelos, os hiperparâmetros foram ajustados por meio da técnica de otimização bayesiana, utilizando a biblioteca python *skopt.gp_minimize* [19]. As Tabelas II, III e IV apresentam uma descrição dos hiperparâmetros que foram otimizados para cada uma das arquiteturas.

TABLE II
HIPERPARÂMETROS DO *Random Forest*.

Hiperparâmetro	Descrição
n_estimators	Nº de árvores
max_depth	Profundidade
min_samples_split	Nº Mín. amostras no nó
min_samples_leaf	Nº Mín. amostras na folha
bootstrap	Descarte de nós

E. Treinamento, avaliação e validação

No treinamento do modelo, utilizou-se uma validação simples devido ao número limitado de amostras, com 70% dos dados destinados ao treino e 30% para teste. A performance

TABLE III
HIPERPARÂMETROS DO *Gradient Boosting*.

Hiperparâmetro	Descrição
n_estimators	Nº estágios do boosting
learning_rate	Taxa de aprendizado
max_depth	Profundidade
min_samples_split	Nº Mín. amostras no nó
min_samples_leaf	Nº Mín. amostras na folha
subsample	Amostras por árvore

TABLE IV
HIPERPARÂMETROS DA *Deep Neural Network*.

Hiperparâmetro	Descrição
num_layers	Nº camadas ocultas
num_units	Nº neurônios por camada
learning_rate	Taxa de aprendizado
dropout_rate	Taxa de dropout

durante o treinamento foi mensurada com a métrica de Raiz do Erro Quadrático Médio (RMSE). Na validação, escolha final do modelo, também foram utilizadas as métricas de Erro Absoluto Médio (MAE) e o coeficiente de determinação (R^2), com o objetivo de aferir o desempenho sob diferentes perspectivas e identificar possíveis ajustes necessários.

Destaca-se que as observações não foram distribuídas de forma equitativa: 66% das amostras (16) estão concentradas entre 90% e 100% de concentração de etanol, conforme ilustrado na Figura 3. Esse desbalanceamento nas observações pode prejudicar o treinamento do modelo.

Para contornar essa questão, o treinamento e a validação foram divididos em duas etapas. Na primeira etapa, as amostras com concentrações entre 90% e 100% foram excluídas, permitindo que o treinamento fosse realizado apenas com 10 amostras equidistribuídas entre 10% e 100%. Na segunda etapa, o treinamento e a validação foram realizados com todas as amostras disponíveis.

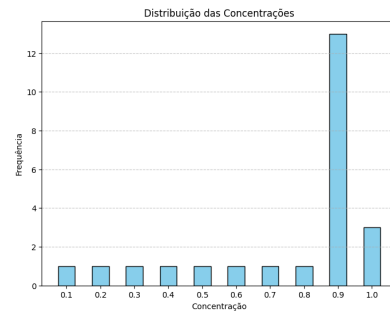


Fig. 3. Histograma da distribuição das amostras de concentração de etanol.

III. RESULTADOS

Em ambos os testes realizados, o desempenho da DNN foi insatisfatório, como mostrado na Figura 4. Esse resultado era esperado, considerando que modelos de redes neurais profundas dependem fortemente de um grande volume de dados para aprender de forma eficaz e capturar com precisão as características da população. Em contextos de dados limitados, a DNN tende a apresentar baixa generalização e pode superestimar ou subestimar padrões devido à insuficiência de exemplos representativos [20].

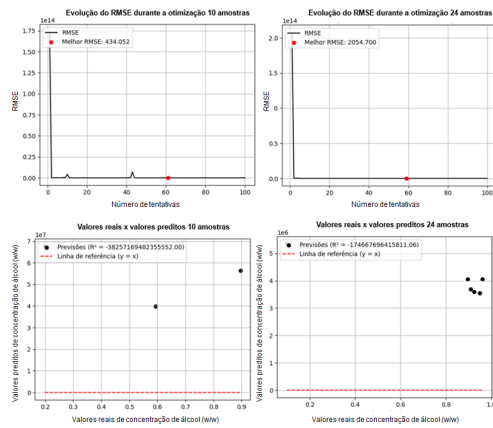


Fig. 4. Resultados do modelo de *Deep Neural Network*.

A. Treinamento com 10 amostras

Utilizando o método de otimização bayesiana do Python, os parâmetros do modelo *Random Forest* foram ajustados conforme mostrado na Tabela V. Nesta etapa, definiu-se os limites de busca das variáveis, ajustando-os sempre que ocorria saturação, para obter o melhor desempenho do modelo. O modelo resultante apresentou baixos valores de erro RMSE e MAE, além de um R^2 elevado (Tabela VI), indicando boa precisão e uma capacidade significativa de explicar a variabilidade dos dados, especialmente para um conjunto de dados reduzido. O R^2 de 0,961 sugere que o modelo captura adequadamente a relação entre as variáveis e o teor de álcool.

TABLE V
SINTONIA DE HIPERPARÂMETROS DO *Random Forest* PARA 10 AMOSTRAS.

Hiperparâmetro	Valor
n_estimators	199
max_depth	500
min_samples_split	2
min_samples_leaf	2
bootstrap	false

TABLE VI
DESEMPENHO DO MODELO *RANDOM FOREST* COM 10 AMOSTRAS.

<i>Random Forest</i>	Valor
RMSE	0.0564
MAE	0.0559
R^2	0.961

A Figura 5, ilustra a relação entre os valores preditos e os valores reais na etapa de validação.

Embora o desempenho do modelo *Gradient Boosting* seja ligeiramente inferior ao do *Random Forest*, ele ainda demonstra boa precisão e capacidade explicativa, com um R^2 de 0,885. O erro um pouco mais alto sugere uma leve desvantagem em relação ao *Random Forest* para o conjunto de amostras. A Tabela VII apresenta os parâmetros otimizados para o modelo, enquanto a Tabela VIII contém os dados de desempenho relacionados à regressão. A Figura 6 ilustra a relação entre os valores preditos e reais na etapa de validação.

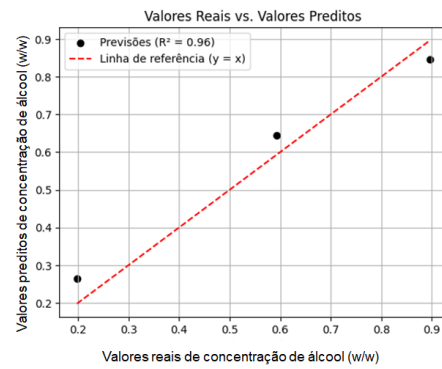


Fig. 5. Resultados do modelo *Random Florest* para 10 amostras.

TABLE VII
SINTONIA DE HIPERPARÂMETROS DO *Gradient Boosting* PARA 10 AMOSTRAS.

Hiperparâmetro	Valor
n_estimators	917
learning_rate	0.067
max_depth	10
min_samples_split	3
min_samples_leaf	2
subsample	0.885

TABLE VIII
DESEMPENHO DO MODELO *Gradient Boosting* COM 10 AMOSTRAS.

<i>Gradient Boosting</i>	
RMSE	0.0970
MAE	0.0822
R^2	0.885

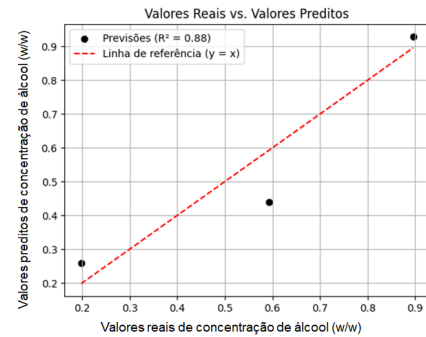


Fig. 6. Resultados do modelo *Gradient Boosting* para 10 amostras.

B. Treinamento com 24 amostras

O modelo *Random Forest* apresentou um desempenho significativamente inferior, conforme indicado na Tabela IX, com um aumento notável nos erros RMSE e MAE e um valor de R^2 de apenas 0,077, evidenciando o ajuste insatisfatório, como mostrado na Figura 7. Esse baixo valor de R^2 indica que o modelo está subajustado, sugerindo que ele não conseguiu capturar a relação entre as variáveis de forma eficaz. Essa limitação pode estar associada a falta de amostras e a distribuição ao longo da faixa de concentração, o que prejudica a capacidade do modelo de generalizar. A Tabela X apresenta os hiperparâmetros ajustados na etapa de treinamento.

TABLE IX
DESEMPENHO DO MODELO *Random Forest* COM 24 AMOSTRAS.

<i>Random Forest</i>	
RMSE	0.2977
MAE	0.1814
R ²	0.077

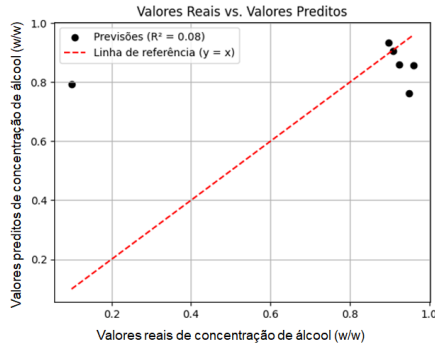


Fig. 7. Resultados do modelo *Random Forest* para 24 amostras.

TABLE X
SINTONIA DE HIPERPARÂMETROS DO *Random Forest* PARA 24 AMOSTRAS.

Hiperparâmetro	Valor
n_estimators	55
max_depth	101
min_samples_split	5
min_samples_leaf	2
bootstrap	true

O modelo *Gradient Boosting* apresentou um desempenho consideravelmente superior ao *Random Forest*, com valores de RMSE e MAE reduzidos e um coeficiente de determinação R^2 de 0,693, conforme mostrado na Tabela XI e Figura 8. Esse resultado indica um ajuste mais robusto, embora ainda aquém da qualidade observada quando o modelo é treinado com um conjunto de 10 amostras. A arquitetura de *Gradient Boosting* mostrou-se mais adequada para lidar com o aumento no número de amostras e o desbalanceamento nas concentrações, o que é consistente com a complexidade desse modelo. Esse tipo de modelo combina múltiplas árvores de decisão em sequência, permitindo que ele corrija erros de predição feitos por árvores anteriores, o que contribui para uma melhor generalização e uma maior robustez ao desbalanceamento dos dados [21].

TABLE XI
DESEMPENHO DO MODELO *Gradient Boosting* COM 24 AMOSTRAS.

<i>Gradient Boosting</i>	
RMSE	0.1715
MAE	0.1218
R ²	0.693

No entanto, o desbalanceamento entre as classes ainda impactou negativamente o desempenho, evidenciando que, embora o *Gradient Boosting* ofereça uma maior resiliência a variações na distribuição de dados, ele ainda pode ser influenciado por uma representação desigual nas classes de concentração. A Tabela XII apresenta os hiperparâmetros ajustados na fase de treinamento.

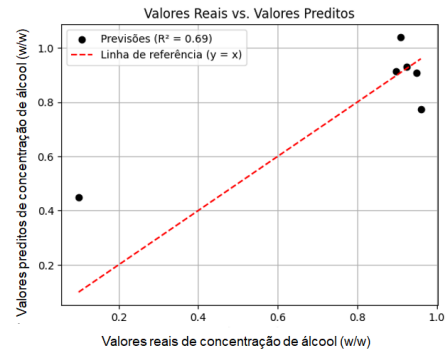


Fig. 8. Resultados do modelo *Gradient Boosting* para 24 amostras.

TABLE XII
SINTONIA DE HIPERPARÂMETROS DO *Gradient Boosting* PARA 24 AMOSTRAS.

Hiperparâmetro	Valor
n_estimators	288
learning_rate	0.096
max_depth	8
min_samples_split	10
min_samples_leaf	2
subsample	0.935

IV. CONCLUSÕES

Na avaliação inicial com 10 amostras, o modelo *Random Forest* apresentou melhor desempenho, possivelmente devido à sua maior capacidade de capturar a variabilidade de um conjunto de dados limitado, além de ser uma abordagem relativamente simples. O *Gradient Boosting* também demonstrou eficácia, mas sofreu uma leve queda de precisão, provavelmente em razão do número reduzido de amostras, insuficiente para um treinamento mais aprofundado.

Ao considerar 24 amostras, o *Gradient Boosting* superou o *Random Forest*, possivelmente por lidar melhor com dados desbalanceados e se ajustar mais adequadamente à concentração de amostras entre 90% e 100% de etanol, o modelo de *ensemble* é mais estável. Em contrapartida, o *Random Forest* parece ter sido prejudicado pela menor diversidade de amostras disponíveis.

De modo geral, considerando este estudo piloto de viabilidade, ambos os modelos demonstraram capacidade para identificar contaminações em etanol, evidenciando seu potencial de aplicação. Esse resultado é especialmente relevante, dado o número limitado de amostras disponíveis e o desbalanceamento presente no conjunto de dados. Entretanto, os resultados ainda não atingem o nível de precisão exigido pela norma [22], que estabelece a detecção de concentrações de etanol abaixo de 99,3% com erro inferior a 0,7%. O *Random Forest* obteve um erro médio de 18,14%, enquanto o *Gradient Boosting* alcançou 12,18%. Para aprimorar esses índices, futuros trabalhos devem focar na ampliação do conjunto de dados em intervalos menores de concentração, na experimentação de outras técnicas de redução de dimensionalidade como alternativas ao PCA, além de investigar o uso de outros modelos que possam apresentar desempenho superior nesse tipo de problema com um número pequeno de amostras.

REFERÊNCIAS

- [1] Guimarães, J.F., “Controle Avançado: Aplicações bem sucedidas são possíveis sim!” 2002.
- [2] Soul Science, “Determinação de teor de Água: O papel do método karl fischer em diversos setores,” 2024, disponível em https://soulsience.com.br/determinacao-de-teor-de-agua-o-papel-do-metodo-karl-fischer-em-diversos_setores/. [Online]. Available: https://soulsience.com.br/determinacao-de-teor-de-agua-o-papel-do-metodo-karl-fischer-em-diversos_setores/
- [3] I. Abuhav, *ISO 9001: 2015-A complete guide to quality management systems*. CRC press, 2017.
- [4] G. R. C. Salles, “Análise de pesquisas sobre a pervaporação como método de separação de etanol,” 2024.
- [5] A. Aydın, C. Keskinoglu, U. Kökbaş, and A. Tuli, “Noninvasive determination of the amount of ethanol in liquid mixtures by ultrasound using bilinear interpolation method,” *Journal of the Mexican Chemical Society*, vol. 64, no. 1, pp. 41–52, 2020.
- [6] Y. Jbari and S. Abderafi, “Liquid density prediction of ethanol/water, using artificial neural network,” *Biointerface Recsearch in Applied Chemistry*, vol. 12, pp. 5625–5637, 2021.
- [7] A. Bernardi, L. E. B. Da Silva, G. F. Veloso, and J. A. Ferreira Filho, “Characterization of the ethanol-water blend by acoustic signature analysis in ultrasonic signals,” *IEEE Access*, vol. 10, pp. 6580–6591, 2022.
- [8] Y. Chen, X. Guo, Y. Pan, Y. Xia, and Y. Yuan, “Dynamic feature splicing for few-shot rare disease diagnosis,” *Medical Image Analysis*, vol. 90, p. 102959, 2023.
- [9] X. Jia, Y. Xiong, J. Zhang, Y. Zhang, and Y. Zhu, “Few-shot radiology report generation for rare diseases,” in *2020 IEEE International Conference on Bioinformatics and Biomedicine (BIBM)*. IEEE, 2020, pp. 601–608.
- [10] Z. Zhang, Z. Zou, N. Li, and Y. Chen, “Classifying galaxy morphologies with few-shot learning,” *Research in Astronomy and Astrophysics*, vol. 22, no. 5, p. 055002, 2022.
- [11] A. Brunel, J. Pasquet, J. Pasquet, N. Rodriguez, F. Comby, D. Fouchez, and M. Chaumont, “A cnn adapted to time series for the classification of supernovae,” *arXiv preprint arXiv:1901.00461*, 2019.
- [12] H. Chen, S. Lindshield, P. I. Ndiaye, Y. H. Ndiaye, J. D. Pruetz, and A. R. Reibman, “Applying few-shot learning for in-the-wild camera-trap species classification,” *AI*, vol. 4, no. 3, pp. 574–597, 2023.
- [13] F. M. Siraj, S. T. K. Ayon, and J. Uddin, “A few-shot learning based fault diagnosis model using sensors data from industrial machineries,” *Vibration*, vol. 6, no. 4, pp. 1004–1029, 2023.
- [14] M. M. Tiago, “Desenvolvimento de uma célula para medição de propriedades de líquidos por ultrassom com manipulação de amostras através de cubetas,” Ph.D. dissertation, 2018.
- [15] ORGANISATION INTERNATIONALE DE MÉTROLOGIE LÉGALE- OIML, “OIML bulletin n. 3,” OIML, Paris, Tech. Rep., 2015.
- [16] A. Natekin and A. Knoll, “Gradient boosting machines, a tutorial,” *Frontiers in neurorobotics*, vol. 7, p. 21, 2013.
- [17] R. Hansch, *Handbook of random forests: Theory and applications for remote sensing*. World Scientific Publishing Co Pte Ltd, 2018.
- [18] P. Nakamoto, *Neural Networks and Deep Learning: Deep Learning explained to your granny A visual introduction for beginners who want to make their own Deep Learning Neural Network*. CreateSpace Independent Publishing Platform, 2017.
- [19] Scikit-Optimize, “SCIKIT-OPTIMIZE Documentation: Release 0.8.1,” Sep. 2020, disponível em https://scikit-optimize.github.io/0.8/_downloads/scikit-optimize-docs.pdf. [Online]. Available: https://scikit-optimize.github.io/0.8/_downloads/scikit-optimize-docs.pdf
- [20] P. R. A. S. Bassi and R. R. de Faissol Attux, “Fundamentos de redes neurais profundas,” *Revista dos Trabalhos de Iniciação Científica da UNICAMP*, no. 26, 2019.
- [21] C. Bentéjac, A. Csörgő, and G. Martínez-Muñoz, “A comparative analysis of gradient boosting algorithms,” *Artificial Intelligence Review*, vol. 54, pp. 1937–1967, 2021.
- [22] “Resolução anp n° 828 de 01/09/2020,” 2020, disponível em: <https://www.normasbrasil.com.br/norma/?id=422148>. Acesso em: 01 nov. 2024.